

Governing Artificial Intelligence: Rule of Law, Democratic Accountability and the Politics of Regulation in a Fragmented World

Monika, Research Scholar, Department of Political Science, Faculty of Humanities, Baba Mastnath University, Rohtak, Haryana, India

Dr. Ravinder Kumar, Associate Professor of Law, Sri Sukhmani College of Law, Dera Bassi, Punjab, India

Abstract

Artificial intelligence has moved from a technical curiosity to a central force in economic, social and political life. As governments race to regulate it, a patchwork of laws and voluntary codes has emerged, led by the European Union's AI Act, the United States' mixture of federal caution and state-level activism, and Asian approaches ranging from Singapore's agentic AI framework to China's algorithm-specific rules. This review paper uses a socio-legal method to examine whether these frameworks genuinely protect the rule of law, democratic accountability and fundamental rights, or whether they mainly manage economic competition between trading blocs. The paper argues that current AI governance suffers from three connected problems: a widening gap between binding law and fast-moving technology, a concentration of regulatory capacity in a small number of wealthy states and firms, and a weak link between formal legal compliance and genuine democratic control over automated power. Drawing on recent regulatory developments up to mid-2026, including the EU's Digital Omnibus and the rise of autonomous AI agents, the paper offers a critical assessment of enforcement gaps and proposes a rights-centred, participatory model of governance that treats AI regulation as a question of political power, not only of technical risk management.

Keywords: *Artificial intelligence governance; rule of law; EU AI Act; algorithmic accountability; digital sovereignty; socio-legal studies; democratic governance; regulatory fragmentation*

1. Introduction

Artificial intelligence (AI) is no longer a subject confined to computer science departments and technology firms. It now shapes who gets a bank loan, which job applications are shortlisted, how police forces allocate patrols, and what information citizens see when they open a news application. Because AI systems influence decisions that were once the exclusive domain of human officials, courts and elected bodies, the question of how AI is governed has become a core question of law and politics, not merely of engineering.

The years 2024 to 2026 mark a turning point in this story. The European Union's Artificial Intelligence Act, the first comprehensive statute of its kind anywhere in the world, entered into force in August 2024 and became fully applicable from August 2026, though several of its high-risk provisions have since been softened and delayed through a 'Digital Omnibus' simplification package agreed in 2026. In the United States, the absence of a comprehensive federal AI statute has pushed individual states such as Colorado and Texas to legislate on algorithmic discrimination and transparency, creating a fragmented domestic landscape even before international differences are considered. Singapore has gone further than most jurisdictions by publishing, in January 2026, the first governance framework aimed specifically at autonomous AI agents, systems that do not merely recommend actions but actually execute them in the world. China continues to regulate AI through a series of targeted administrative rules focused on algorithmic recommendation, deep synthesis and generative content, reflecting a state-centric model of control quite different from the EU's rights-based approach.

This proliferation of rules might appear, at first glance, to be good news for the rule of law: more law generally implies more predictability and more accountability. Yet a closer, socio-legal reading suggests a more complicated picture. Laws are not created in a political vacuum. They are shaped by lobbying, by competition between trading blocs for technological advantage, and by the practical difficulty of regulating systems that change faster than legislative cycles. The central question this paper asks is therefore not simply 'is AI regulated?' but 'whose interests does AI regulation actually serve, and does it strengthen or weaken democratic and legal accountability over automated power?'

This review paper contributes to the growing socio-legal literature on technology governance by offering a comparative and critical reading of the current regulatory landscape. It treats AI governance as a site of political contest between states, corporations and citizens, rather than as a purely technical exercise in risk classification. The paper is organised as follows. Section two sets out the objectives and method. Section three reviews the existing literature on law, technology and democratic accountability. Section four maps the fragmented global regulatory landscape as it stood in mid-2026. Section five examines the socio-political dimensions of AI governance, including concerns about power concentration and digital inequality. Section six turns to legal and constitutional challenges, including the difficulty of applying traditional doctrines of accountability to autonomous systems. Section seven offers a critical analysis of the tensions within current frameworks. Section eight proposes a set of reform directions, and section nine concludes.

2. Objectives and Methodology

The paper has three objectives. First, to map the major AI governance frameworks that were operative or under active development as of mid-2026, with attention to their legal form, scope and enforcement mechanisms. Second, to critically evaluate these frameworks against long-standing socio-legal benchmarks of the rule of law, namely predictability, accountability, proportionality and access to remedy. Third, to identify the structural gaps that remain,

particularly around autonomous AI agents, cross-border enforcement and the unequal capacity of different states to regulate powerful technology firms.

This is a doctrinal and socio-legal review paper. It does not present new empirical data collected through surveys or interviews. Instead, it synthesises secondary material, including primary legal texts such as the EU AI Act and its amending Omnibus package, regulatory guidance from bodies such as the European Commission's AI Office, and analytical commentary from legal scholars, policy think tanks and international organisations including the Organisation for Economic Co-operation and Development. The socio-legal method used here treats law not as a closed technical system but as one social institution among others, shaped by economic interests, institutional capacity and political ideology. This method is well suited to a review of AI governance because it allows the paper to ask not only what the law says, but who benefits from it and who is left without effective protection.

A limitation of this method should be acknowledged at the outset. Because AI regulation is changing rapidly, with the EU's Digital Omnibus negotiations and several national frameworks still being finalised at the time of writing, some of the specific details discussed here may be superseded by later amendments. The paper therefore focuses on structural and institutional patterns that are likely to persist even as individual rules change, rather than on provisions that may soon be revised.

3. Literature Review

The scholarly literature on technology governance has long grappled with what is sometimes called the 'pacing problem': the tendency of technology to develop faster than the law can respond. Early work on internet governance in the 1990s and 2000s already noted that legislatures struggle to write rules general enough to survive rapid technical change, yet specific enough to be enforceable. This tension has resurfaced with even greater intensity in the context of AI, where systems can be retrained, redeployed or repurposed within weeks of a law being passed.

A second strand of literature concerns algorithmic accountability. Scholars working in this tradition have argued that automated decision-making challenges classical administrative law doctrines, which typically assume a human decision-maker who can explain the reasons for a decision and be held answerable for it. When decisions are produced by complex statistical models, the traditional chain of accountability, from official to institution to citizen, becomes harder to trace. This has led to proposals for new legal concepts such as a 'right to explanation', algorithmic impact assessments and mandatory human oversight, several of which have since been incorporated, in varying degrees, into the EU AI Act and comparable instruments elsewhere.

A third body of work, closer to political economy, treats AI regulation as an instrument of geopolitical and economic competition rather than a neutral exercise in consumer protection. On this view, the EU's early and comprehensive legislative move can be read as an attempt to

set a global standard in the way its General Data Protection Regulation did for data privacy, a phenomenon sometimes called the 'Brussels effect', where multinational firms adopt the strictest jurisdiction's rules everywhere to reduce compliance costs. The United States, by contrast, has generally favoured lighter-touch, innovation-first policies at the federal level, reflecting both a different constitutional culture around speech and commerce and the economic interest of dominant American technology firms in avoiding heavy compliance burdens. China's regulatory model, meanwhile, is frequently analysed as an instrument of state control over information flows and social stability rather than as a rights-protective measure in the European sense.

A fourth and more recent strand of scholarship, directly relevant to this paper, addresses the governance of autonomous or 'agentic' AI systems, meaning models that can plan multi-step tasks and take actions in digital or physical environments without continuous human instruction. Commentators have observed that most existing frameworks, including the EU AI Act, were substantially negotiated before the widespread emergence of such agents and therefore assume a model in which AI assists a human decision, rather than one in which AI acts largely on its own. This gap is treated by several analysts as the most urgent unresolved problem in AI governance during 2026.

Finally, a smaller but growing literature situates AI governance within broader debates on digital colonialism and global inequality, arguing that the technical and financial capacity to build, audit and regulate advanced AI systems remains concentrated in a handful of wealthy states and corporations, leaving lower-income countries as rule-takers rather than rule-makers. This paper draws on all four strands to build an integrated socio-legal critique, rather than treating AI governance as a purely technical or purely economic subject.

4. The Fragmented Global Landscape of AI Regulation

As of mid-2026, no single global treaty governs artificial intelligence. Instead, a patchwork of binding statutes, administrative regulations and voluntary codes of conduct has emerged, differing sharply in legal force, scope and philosophy. This section reviews the major approaches before turning, in later sections, to their socio-political and legal implications.

4.1 The European Union: A Risk-Based, Rights-Oriented Model

The European Union's AI Act remains the most comprehensive binding AI statute in the world. It classifies AI systems into four risk tiers, unacceptable, high, limited and minimal, banning practices such as government-run social scoring and certain forms of manipulative or exploitative AI outright, while subjecting high-risk systems, such as tools used in recruitment, credit scoring, law enforcement or education, to obligations including risk management, technical documentation, human oversight, data governance and post-market monitoring. The general provisions on prohibited practices and AI literacy took effect in February 2025, governance rules and obligations for general-purpose AI models became applicable in August

2025, and the remaining core provisions, including most transparency obligations, became applicable from August 2026.

However, the Act's practical bite has been softened considerably by a 'Digital Omnibus' simplification package, on which EU institutions reached political agreement in May 2026. This package extends compliance deadlines for high-risk systems, particularly those embedded in regulated products already governed by the EU's Machinery Regulation, narrows the definition of 'safety component' so that many everyday AI applications fall outside the high-risk category, and extends simplified compliance pathways originally designed for small and medium enterprises to mid-sized companies as well. National-level regulatory sandboxes, meant to let smaller innovators test AI systems under supervision, have also been postponed by a year, from August 2026 to August 2027. Read together, these amendments illustrate a recurring pattern in technology regulation: ambitious rights-based legislation is frequently diluted during the implementation phase in response to industry lobbying and concerns about competitiveness, even where the underlying social harms the legislation was designed to address remain unchanged.

At the same time, the amended Act has added new prohibitions that respond to public concern rather than industrial lobbying, most notably a ban, effective December 2026, on using AI to generate non-consensual intimate imagery and child sexual abuse material. This shows that the Act is not simply being weakened across the board; rather, its stringency is being reallocated, tightened where public outrage is strongest and loosened where industrial cost pressure is strongest, a pattern that itself deserves closer socio-legal attention because it reveals how legislative compromise, rather than technical necessity, determines the final shape of supposedly neutral risk categories.

4.2 The United States: Federal Restraint, State Activism

The United States has not adopted a comprehensive federal AI statute comparable to the EU AI Act. Federal policy has instead relied on executive orders and voluntary frameworks, most notably the National Institute of Standards and Technology's AI Risk Management Framework, which offers guidance rather than binding rules. In the resulting vacuum, individual states have moved to legislate. Colorado's AI Act targets algorithmic discrimination in high-risk systems and imposes documentation and transparency duties on developers and deployers, while Texas has introduced its own Responsible Artificial Intelligence Governance Act. This has produced a genuinely fragmented domestic landscape, where a single AI system may be subject to different disclosure and fairness obligations depending on which state its users are located in, and where an ongoing federal effort to preempt state AI laws remains legally unsettled.

This American model reflects a distinct political philosophy: a preference for market-led innovation, a constitutional tradition protective of commercial and expressive freedom, and a federal structure that allows states to act as policy laboratories. From a rule-of-law perspective, the American approach offers flexibility and room for experimentation, but at the cost of

predictability for businesses and, more importantly, inconsistent protection for citizens depending on their state of residence.

4.3 Asian Approaches: Singapore's Agentic Framework and China's Administrative Model

Singapore's Infocomm Media Development Authority released, in January 2026, what is widely regarded as the first governance framework designed specifically for autonomous AI agents. The framework introduces a graduated five-tier taxonomy of autonomy, from tool-assisted systems to fully autonomous agents, with governance obligations increasing at each tier, alongside a proposed 'Agent Identity Card' that would require AI agents to disclose their capabilities, limitations and authorised scope of action, and a framework allocating liability between the entity that builds an agent platform and the entity that deploys it for a specific task. This framework is significant because it directly addresses the governance gap identified in the literature review: most existing rules assume AI recommends rather than acts, and Singapore's model is among the first serious legal attempts to close that gap.

China, by contrast, regulates AI through a series of targeted administrative measures covering algorithmic recommendation systems, deep synthesis technology and generative AI services, requiring registration, security assessments and content controls. This model prioritises state oversight of information flows and social stability over the individual-rights framing found in the EU Act, illustrating that AI governance approaches are not simply more or less strict versions of the same idea, but reflect different underlying theories of what the law is for, protecting individual autonomy in the European model, or preserving state authority and social order in the Chinese model.

4.4 The Global South and the Question of Regulatory Capacity

Much of the comparative literature on AI governance focuses on the EU, the United States and East Asia, leaving the position of lower- and middle-income countries under-examined. Yet these countries face a distinct set of problems. Few have the technical capacity to conduct independent audits of advanced AI models, and most lack the market size to compel multinational technology firms to comply with locally-drafted rules in the way the EU can. As a result, countries in Africa, South Asia and Latin America are often placed in the position of importing regulatory templates drafted elsewhere, whether the EU's rights-based model or lighter voluntary codes favoured by American firms, rather than designing frameworks suited to their own social and economic priorities. India, for instance, has so far relied on a mixture of existing data protection law, sector-specific guidance from its telecommunications and electronics ministry, and voluntary advisories to platforms, rather than a dedicated AI statute, reflecting a cautious, incremental approach common among large developing economies that wish to avoid discouraging domestic investment while still signalling concern about algorithmic harm.

5. Socio-Political Dimensions of AI Governance

Law does not operate in isolation from social and political structures, and AI governance illustrates this especially clearly. This section examines three socio-political dimensions that a purely doctrinal reading of AI statutes tends to overlook: the concentration of power, the uneven distribution of harm and benefit, and the erosion of democratic input into decisions that increasingly govern everyday life.

5.1 Concentration of Technical and Political Power

Advanced AI systems require enormous computing infrastructure, vast quantities of data and highly specialised technical talent, resources concentrated in a small number of firms based mainly in the United States and China. This concentration has political consequences that go beyond ordinary market competition. When only a handful of companies can build and audit the most capable AI systems, regulators become structurally dependent on those same companies for the technical expertise needed to write and enforce rules. The European Union's own AI Office, for example, is required to build in-house evaluation capacity for advanced models, a task that is expected to remain limited until at least 2027, meaning that for the near term, effective oversight of the most powerful systems continues to depend heavily on information voluntarily disclosed by the firms being regulated. This dynamic, sometimes called 'regulatory capture through complexity', is not unique to AI, but the pace and technical opacity of AI development make it especially pronounced here.

A further political consequence concerns the relationship between AI governance and interstate competition. Regulatory frameworks are increasingly framed, in political rhetoric if not always in legal text, as instruments of economic and strategic competitiveness rather than purely as protections for citizens. The EU's willingness to delay and soften high-risk obligations in 2026, discussed above, was explicitly justified in policy documents by concerns that overly strict rules would put European firms at a disadvantage relative to American and Chinese competitors. This suggests that even in jurisdictions with strong rights-based traditions, the political weight given to industrial competitiveness can outstrip the political weight given to the citizens the legislation was originally designed to protect.

5.2 Unequal Distribution of Benefit and Harm

The social groups most likely to be affected by high-risk AI applications, such as automated hiring tools, predictive policing systems and welfare fraud detection algorithms, are frequently the same groups with the least political and economic power to contest adverse decisions. Existing research on algorithmic discrimination has repeatedly found that automated systems trained on historical data can reproduce or even amplify existing patterns of bias against particular racial, gender or socio-economic groups, precisely because that historical data reflects past discriminatory practice. Where legal remedies for such harms depend on technically demonstrating how a model reached its decision, individuals from marginalised

communities, who often have the least access to legal and technical expertise, are placed at a structural disadvantage in seeking redress, even where formal legal rights, such as Colorado's protections against algorithmic discrimination, exist on paper.

This unequal distribution of harm is compounded by an unequal distribution of benefit. Productivity gains from AI adoption have so far accrued disproportionately to capital owners and highly skilled workers, while routine cognitive and administrative jobs, disproportionately held by women and lower-income workers in many economies, face displacement risk. A socio-legal reading of AI governance must therefore ask not only whether a given AI system is technically 'fair' according to a statistical definition, but whether the broader legal and economic structure surrounding its deployment distributes the resulting gains and losses in a way that a democratic society would consider just.

5.3 Democratic Deficit in Standard-Setting

Much of the substantive content of AI regulation is not decided by elected legislatures but by technical standard-setting bodies, industry codes of practice and administrative guidance issued by specialised agencies. The EU AI Act itself explicitly relies on harmonised technical standards, developed by European standardisation bodies, to give operational meaning to broad legal obligations such as 'appropriate human oversight' or 'robust data governance'. While this is a common and often necessary feature of technical regulation, it raises a genuine democratic concern: the bodies drafting these standards are typically composed of industry engineers and technical experts, with limited direct representation from civil society, affected communities or elected officials. The practical content of citizens' rights under AI law is therefore often settled in technical committees rather than in parliamentary debate, a pattern that mirrors long-standing critiques of technocratic governance in fields such as financial regulation and environmental standard-setting, but which is arguably more consequential here because AI systems increasingly mediate access to core civic and economic opportunities.

6. Legal and Constitutional Challenges

Beyond its socio-political dimensions, AI governance raises distinct legal and constitutional problems that test the adequacy of existing doctrine. This section considers four such problems: the erosion of traditional accountability chains, the enforcement gap between law on paper and law in practice, the specific challenge posed by autonomous AI agents, and the tension between AI regulation and other constitutional values such as free expression and privacy.

6.1 The Erosion of Traditional Accountability Chains

Administrative and constitutional law has traditionally relied on identifiable human decision-makers who can be asked to justify a decision and can be held liable if that decision is unlawful. When a bank refuses a loan application, or a public authority denies a benefit claim, on the basis of a statistical model's output, this chain of accountability becomes difficult to trace. It is often unclear whether responsibility lies with the developer who built the underlying model,

the deployer who integrated it into a specific decision-making process, or the human official who formally approved the output without meaningfully reviewing it, sometimes referred to in the literature as 'rubber-stamping'. The EU AI Act attempts to address this by imposing distinct obligations on 'providers' and 'deployers' of high-risk systems, and by requiring 'appropriate human oversight', but critics have observed that human oversight requirements can become a symbolic formality if the human reviewer lacks the time, training or authority to meaningfully depart from the AI system's recommendation. Without stronger procedural safeguards, such as a genuine and documented opportunity to override automated outputs, the rule-of-law value of a right to reasoned decision-making risks becoming a paper guarantee rather than a practical one.

6.2 The Enforcement Gap

A statute is only as strong as the institutions that enforce it. The EU AI Act provides for substantial penalties, up to fifteen million euros or three percent of global annual turnover for violations of transparency obligations, and considerably higher for the most serious prohibited practices. Yet penalties on paper do not guarantee enforcement in practice. Enforcement of the AI Act is distributed across national market surveillance authorities in twenty-seven member states, alongside a central AI Office responsible for the most powerful general-purpose models, creating a coordination challenge similar to that already observed under the EU's General Data Protection Regulation, where enforcement intensity has varied considerably between member states depending on the resources and political will of national data protection authorities. Comparable enforcement gaps exist in the American state-law context, where individual state attorneys general and administrative agencies, often working with limited budgets, are expected to police compliance by some of the best-resourced legal departments in the world.

This enforcement gap is not merely an administrative inconvenience; it has a direct rule-of-law dimension. A legal system in which rights exist formally but are rarely enforced in practice risks producing what socio-legal scholars have long called 'law in the books' versus 'law in action', a divergence that tends to erode public trust in law generally, not only in the specific rules concerned.

6.3 The Agentic AI Gap

The most significant emerging legal challenge concerns autonomous AI agents, systems capable of planning and executing multi-step tasks, such as making purchases, sending communications, modifying code or interacting with other software systems, with limited or no continuous human instruction. Most existing legal frameworks, including the EU AI Act, were substantially negotiated before such agents became commercially widespread, and their risk categories generally assume a model that assists a human decision rather than one that independently takes action. Singapore's January 2026 framework represents an early attempt to close this gap through its graduated autonomy taxonomy and its proposed allocation of liability between platform builders and deployers, but this remains a voluntary governance

framework rather than binding law, and no jurisdiction has yet enacted comprehensive, enforceable rules specifically addressing agentic AI.

From a legal standpoint, agentic systems strain foundational doctrines of causation and foreseeability that underpin tort and criminal liability. If an autonomous agent, acting on a broadly worded instruction, causes harm through an action its human principal did not specifically foresee or authorise, existing liability rules struggle to determine whether responsibility should fall on the person who issued the instruction, the company that built the underlying model, or the entity that deployed the agent in a particular context. Until legislatures directly address this question, courts in different jurisdictions are likely to reach inconsistent answers, undermining the predictability that is a core rule-of-law value.

6.4 Tension with Other Constitutional Values

AI regulation does not operate in a vacuum separate from other constitutional guarantees. Transparency obligations requiring disclosure of AI-generated content, for instance, must be balanced against freedom of expression, since overly broad labelling requirements risk chilling legitimate satire, art or political commentary that makes creative use of AI tools. Similarly, obligations that permit processing of sensitive personal data, such as health or biometric information, in order to detect and correct algorithmic bias, as newly permitted under the EU's 2026 Omnibus amendments, create a direct tension with data protection principles that ordinarily restrict the use of such sensitive categories of data. These tensions illustrate that AI governance cannot be treated as a self-contained regulatory silo; it must be read alongside, and sometimes reconciled with, pre-existing constitutional and human rights frameworks, a task that legislatures have so far addressed only partially.

7. Critical Analysis: Tensions and Contradictions in Current Frameworks

Bringing together the socio-political and legal analysis above, this section identifies three underlying contradictions that, in the author's assessment, define the current period of AI governance.

The first contradiction is between the stated aim of protecting fundamental rights and the practical effect of protecting competitiveness. Every major framework examined in this paper, from the EU Act to Singapore's agentic framework to American state legislation, is publicly justified primarily in terms of protecting citizens from harm, whether discrimination, manipulation or unsafe deployment. Yet the actual legislative and administrative choices made, particularly the EU's 2026 decision to delay and narrow high-risk obligations, are frequently justified internally by reference to competitiveness concerns rather than by any change in the underlying evidence of harm. This suggests that AI governance functions, in practice, as a site where rights-based and market-based logics compete, with the balance currently tilting toward the latter whenever the two are perceived to conflict.

The second contradiction is between the global nature of AI systems and the national or regional nature of the law that governs them. A single large AI model may be trained in one jurisdiction, hosted on servers in a second, and deployed to make consequential decisions about individuals in dozens of others. Yet legal accountability remains overwhelmingly territorial, dependent on where a company is headquartered, where its users are located, or where a harmful decision technically occurred. This mismatch allows sophisticated technology firms to engage in a degree of regulatory arbitrage, structuring their operations to minimise exposure to the strictest applicable rules, a dynamic long familiar from other areas of transnational corporate regulation such as tax law, but now extended into a domain, automated decision-making, that touches far more directly on individual rights.

The third contradiction is between the technical framing of AI regulation as a matter of 'risk management' and its actual character as a question of political power. Much of the language used in the EU AI Act and comparable instruments, terms such as 'risk tiers', 'conformity assessment' and 'technical documentation', borrows heavily from product safety regulation, treating an AI system somewhat like a household appliance whose risks can be objectively measured and certified. This framing has real advantages, offering a workable and relatively depoliticised vocabulary for legislators to agree on. But it also obscures a more uncomfortable truth: decisions about which AI applications are permitted, which are restricted and which are banned outright are, at root, decisions about how much automated power over human lives a society is willing to tolerate, and who is entitled to wield it. Treating these as purely technical risk calculations, rather than openly political choices requiring broad democratic deliberation, risks depoliticising questions that properly belong in the public sphere, not solely in the hands of engineers and compliance officers.

Taken together, these three contradictions suggest that the current generation of AI governance frameworks, while a genuine improvement over the near-total absence of binding rules that characterised the early 2020s, remains structurally weighted toward the interests of the states and firms best positioned to shape it. This is not a claim that current regulators are acting in bad faith; rather, it is an observation, consistent with long-standing socio-legal scholarship on regulatory politics, that law reflects the balance of power among the actors involved in drafting it, and that balance currently favours large technology firms and the states that host them over individual citizens, smaller businesses and lower-income countries.

8. Towards a Socio-Legal Framework for Reform

Based on the analysis above, this paper proposes five directions for reform, framed not as a technical checklist but as a set of political and institutional priorities.

First, enforcement capacity should be treated as seriously as legislative drafting. A right that cannot realistically be enforced offers little protection. Regulators, whether the EU's AI Office, national market surveillance authorities, or state attorneys general in the United States, require dedicated, ring-fenced funding and technical staff independent of the firms they regulate, rather than continued reliance on information voluntarily disclosed by those firms.

Second, legal frameworks must be extended, urgently, to address autonomous AI agents directly, rather than relying on categories designed for earlier, recommendation-only systems. This should include clear default liability rules allocating responsibility between instruction-givers, model developers and deployers, along the lines suggested by Singapore's graduated autonomy taxonomy, but placed on a binding statutory footing rather than left as voluntary guidance.

Third, technical standard-setting processes that give operational content to broad legal duties should be opened to meaningful participation by civil society organisations, affected communities and independent academic experts, not left to industry-dominated technical committees alone. Democratic legitimacy requires that the practical meaning of citizens' rights not be settled entirely outside public view.

Fourth, international cooperation mechanisms are needed to reduce regulatory arbitrage, potentially through mutual recognition agreements between jurisdictions with comparable rights protections, combined with continued support for the OECD's AI Policy Observatory and similar bodies that track and compare national frameworks, so that lower-capacity states are not forced to choose between adopting an ill-fitting foreign template wholesale or leaving citizens entirely unprotected.

Fifth, and most fundamentally, AI governance debates should be reframed, in public discourse and in legislative process, as questions of democratic political choice about the appropriate limits of automated power, rather than as narrow technical exercises in risk certification. This reframing would not by itself resolve the underlying power imbalances identified in this paper, but it would make those imbalances more visible, and therefore more open to contestation through ordinary democratic and legal channels, including litigation, parliamentary scrutiny and public advocacy.

9. Conclusion

This paper has examined the rapidly evolving landscape of artificial intelligence governance as it stood in mid-2026, from the European Union's AI Act and its subsequent softening through the Digital Omnibus package, to fragmented state-level regulation in the United States, Singapore's pioneering framework for autonomous agents, and China's state-centred administrative model. Using a socio-legal method, the paper has argued that these frameworks, whatever their differences in legal form and underlying philosophy, share three structural weaknesses: a persistent gap between formal legal obligation and practical enforcement, a concentration of both technical capacity and political influence in a small number of dominant states and firms, and an underlying tendency to frame what are essentially political choices about acceptable automated power as narrow technical questions of risk management.

None of this suggests that the current wave of AI legislation is worthless. The EU AI Act, for all its subsequent dilution, has already shifted global conversation by establishing binding categories of prohibited and high-risk AI practice where none existed before, and Singapore's

agentic AI framework has usefully named a governance gap that other jurisdictions will eventually need to address. But a critical, socio-legal reading cautions against treating the mere existence of AI statutes as evidence that the rule of law has been secured over this powerful new technology. Genuine accountability will require sustained attention not only to what the law says, but to who enforces it, who shapes its technical content, and whose interests it ultimately serves. As autonomous AI systems become more capable and more widely deployed in the years ahead, the gap between formal legal text and lived democratic accountability is likely to widen further unless the reforms proposed in this paper, or comparable measures, are pursued with genuine political commitment.

References

- European Commission, 'Regulatory Framework Proposal on Artificial Intelligence', Shaping Europe's Digital Future (Digital Strategy portal, 2024–2026).
- European Commission, 'Navigating the AI Act: Questions and Answers', Shaping Europe's Digital Future (2026).
- Regulation (EU) 2024/1689 of the European Parliament and of the Council laying down harmonised rules on artificial intelligence (Artificial Intelligence Act), Official Journal of the European Union, 2024.
- Latham & Watkins, 'AI Act Update: EU Resolves to Change Rules and Extend Deadlines', Client Alert (May 2026).
- Global Policy Watch (Covington & Burling), 'EU AI Act Update: Timeline Relief, Targeted Simplification, and New Prohibitions' (June 2026).
- Legal Nodes, 'EU AI Act 2026 Updates: Compliance Requirements and Business Risks' (2026).
- OneTrust, 'Where AI Regulation is Heading in 2026: A Global Outlook' (2026).
- Colibra, 'AI Regulatory Compliance in 2026: EU AI Act, US Orders, and State Laws' (2026).
- Infocomm Media Development Authority of Singapore, 'Model AI Governance Framework for Agentic AI' (January 2026).
- Organisation for Economic Co-operation and Development, 'OECD AI Policy Observatory' (accessed 2026).
- Novelli, C., Hacker, P., Morley, J., Trondal, J., and Floridi, L., 'Robust AI Governance: A Response to Emerging Regulatory Challenges', analysis published via artificialintelligenceact.eu (2025–2026).
- National Institute of Standards and Technology, 'AI Risk Management Framework (AI RMF 1.0)', U.S. Department of Commerce.
- Colorado General Assembly, Colorado Artificial Intelligence Act (SB 24-205).
- Texas Legislature, Texas Responsible Artificial Intelligence Governance Act (effective 1 January 2026).
- Sombra Inc., 'An Ultimate Guide to AI Regulations and Governance in 2026' (2025–2026).
- Pasquale, F., *The Black Box Society: The Secret Algorithms that Control Money and Information* (Harvard University Press, 2015).
- Zuboff, S., *The Age of Surveillance Capitalism* (PublicAffairs, 2019).
- Bradford, A., *The Brussels Effect: How the European Union Rules the World* (Oxford University Press, 2020).